

AI 대화형 에이전트의 시각적 의인화 효과: 과업지향 맥락에서의 TAM 기반 실험연구

정재열¹ · 윤홍석^{2*}

¹한국과학기술정보연구원 연구지능화센터 선임기술원

²건국대학교 모빌리티인문학 연구원 HK 연구교수

Effects of Visual Anthropomorphism in AI Conversational Agents: A TAM-Based Experimental Study in Task-Oriented Contexts

Jaeyeol Jeong¹ · Hongsuk Yoon^{2*}

¹Senior Engineer, Research Data Intelligence Center, Korea Institute of Science and Technology Information (KISTI), Daejeon 34141, Korea

²HK Research Professor, Academy of Mobility Humanities, Konkuk University, Seoul 05029, Korea

[요약]

본 연구는 AI 대화형 에이전트의 시각적 의인화 단서가 과업지향적 환경에서 사용자의 사전 수용 평가에 어떠한 영향을 미치는지를 검증하였다. 195명을 대상으로 2(시각 단서: 인간유사형 vs. 기계형)×2(형태 단서: 전문형 vs. 일반형) 집단 간 실험을 실시하기 위해, 12개의 일상 과업을 담은 영상 시뮬레이션을 제시하였다. TAM에 기반하여 지각된 사용용이성과 유용성의 연속 매개 효과를 분석한 결과, 인간유사형 시각 단서는 기계형 피드백에 비해 지각된 사용용이성을 유의하게 낮추었고, 이 효과는 지각된 유용성을 거쳐 태도와 사용의도의 감소로 이어졌다. 반면 신뢰는 동일한 매개 경로를 따르지 않았으며, 이는 추가적 선행요인이 필요할 가능성을 시사한다. 형태 단서 비교에서는 일반형이 높은 사용용이성을 보였으나, 조작점검의 한계상 이는 일관된 장치 형태가 적응 부담을 낮춘 결과일 수 있어 탐색적 수준에서만 해석한다. 일상적 과업 맥락의 인간-인공지능 상호작용에서 시각 단서를 실용적 과업 요구에 맞게 조정하는 설계가 중요함을 시사한다.

[Abstract]

This study examines whether visual anthropomorphic cues in a screen-based AI conversational agent affect pre-adoption evaluations in task environments. A 2 (visual cue: human-like versus machine-like) × 2 (form-based role cue: specialist versus generalist) between-subjects experiment was conducted with 195 participants using video-based simulations of 12 everyday tasks. Based on the technology acceptance model (TAM), serial mediation through perceived ease of use (PEOU) and perceived usefulness (PU) was tested. Human-like visual cues were associated with significantly lower PEOU than machine-like feedback, and this effect carried through PU to reduce attitude and intention to use through mediation. However, trust did not follow this pathway, suggesting that building trust may require additional antecedents. The findings suggest the importance of adapting visual cues to the pragmatic demands of human-AI interaction, while the form-based role cue findings are interpreted only as exploratory due to limitation in the manipulation check.

색인어 : 인공지능 대화형 에이전트, 시각적 의인화, 지각된 사용용이성, 기술수용모형, 인간-인공지능 상호작용

Keyword : AI Conversational Agents, Visual Anthropomorphism, Perceived Ease of Use, Technology Acceptance Model (TAM), Human-AI Interaction

<http://dx.doi.org/10.9728/dcs.2026.27.6.1627>



This is an Open Access article distributed under the terms of the Creative Commons Attribution Non-Commercial License (<http://creativecommons.org/licenses/by-nc/3.0/>) which permits unrestricted non-commercial use, distribution, and reproduction in any medium, provided the original work is properly cited.

Received 07 April 2026; Revised 30 April 2026

Accepted 22 May 2026

*Corresponding Author; Hongsuk Yoon

Tel: ※ 개인정보 표시제한

E-mail: hongsukyoon@konkuk.ac.kr

I. Introduction

AI conversational agents—such as smart speakers, in-car assistants, and screen-based home devices—have become increasingly embedded in everyday routines. As these systems increasingly function as persistent content interfaces, the question of which design cues support or hinder initial acceptance in routine task settings has become both practically and theoretically important[1],[2]. The Computers Are Social Actors (CASA) paradigm suggests that users tend to respond socially to technologies that display social cues[3], while anthropomorphism theory explains how human-like features invite human-centered inference in technology use[4]. Consistent with these perspectives, prior studies have reported that anthropomorphic cues can enhance social presence, engagement, and related evaluations, although the magnitude of these effects varies across contexts[5]–[7]. However, these effects are not universal. Karimova[8] argued that anthropomorphic cues can inflate expectations beyond what expert chatbots can functionally deliver. Kang and Kim[9] found that humanization increased cognitive load in task-oriented Internet of Things (IoT) settings, and Seeger et al.[10] likewise reported additional cognitive demands in anthropomorphic conversational interfaces. Together, these findings suggest that the effects of anthropomorphic cues on usability are context-sensitive rather than uniformly positive[11]. Despite this growing literature, most prior work has examined anthropomorphism in relational, entertainment, or broad service settings rather than in routine task environments[12]. In addition, although the Technology Acceptance Model (TAM) identifies perceived ease of use (PEOU) and perceived usefulness (PU) as central mediators of technology acceptance[13],[14], few experimental studies have tested whether visual anthropomorphic cues shape downstream evaluations through the PEOU→PU pathway in pre-adoption settings. This study addresses that gap through a 2 (Visual Cue: Human-like vs. Machine-like)×2 (Form-Based Role Cue: Specialist vs. Generalist) between-subjects experiment with 195 participants. They viewed video-based simulations of an AI conversational agent performing 12 everyday tasks across bedroom, vehicle, and living-room contexts. Using TAM as a

mediating framework, the study examines this pathway in a pre-adoption setting.

This study contributes to AI conversational agent acceptance research by shifting attention from relational or entertainment-oriented anthropomorphism to routine task contexts. It shows that a low-fidelity visual anthropomorphic cue may reduce, rather than uniformly improve, initial evaluations by lowering PEOU and, through PU, attitude and intention to use. It also distinguishes this usability-based pathway from trust, which did not follow the same mechanism, while treating the form-based role cue only as an exploratory device-form contrast.

II. Literature Review

2-1 Social Cues and the CASA Framework

The Computers Are Social Actors (CASA) paradigm[3] and Sundar's[1] Theory of Interactive Media Effects (TIME) both suggest that interface cues shape user expectations and evaluations of interactive systems. In AI conversational agents, human-like visual features can activate social responses and influence perceived credibility, usability, and trustworthiness[4],[15]. Prior studies have shown that anthropomorphic cues may enhance perceived usefulness (PU), social presence, and related evaluations[7],[16], but meta-analytic evidence also indicates substantial heterogeneity across contexts[5]. These mixed findings suggest that the effects of social cues are not uniform and may depend on the interaction setting and user goal.

2-2 Visual Anthropomorphic Cues

Visual anthropomorphism refers to the use of human-like visual features such as eyes, facial expressions, or character-like avatars in interface design[4],[5],[12]. Prior research has often reported positive effects of such cues on engagement, warmth, and social presence, particularly in relational or service-oriented contexts[5],[6],[17]. However, recent studies have also shown that anthropomorphic cues can produce neutral or negative effects when they elevate expectations beyond a system's functional capabilities[8]–[10]. This possibility is especially

relevant in utilitarian environments, where users prioritize efficient task completion over social engagement. In such contexts, human-like cues may create an additional interpretive demand and increase cognitive burden rather than support usability[8],[9]. This interpretation is consistent with expectation disconfirmation theory[18], which suggests that when appearance signals greater capability than the system can deliver, users may evaluate the interaction less favorably[19]. The present study focuses specifically on a low-fidelity, screen-based visual anthropomorphic cue: animated eye expressions versus abstract data visualization. Because this manipulation is limited to the visual modality, the findings are bounded to screen-based visual anthropomorphic cues rather than multimodal or high-fidelity anthropomorphism[20]–[22]. Based on these considerations, the following research question is posed:

RQ1: In routine task contexts, how do visual anthropomorphic cues (human-like vs. machine-like) on an AI conversational agent affect users' perceived ease of use, perceived usefulness, trust, attitude, and intention to use?

2-3 Form-Based Role Cues

A second design dimension relevant to conversational agents is role framing: whether an agent appears to be specialized for particular tasks or usable as a general-purpose assistant[1],[11],[13],[23]. In interface design, such role expectations may be conveyed not only through verbal labels or claims of expertise, but also through device form. A task-specific embedded form may imply functional specialization, whereas a consistent standalone form may imply general-purpose availability across contexts[24]–[26]. However, form-based role cues are conceptually distinct from expertise-based specialization. Unlike explicit expertise labels or demonstrated domain competence, device form may influence evaluations through several overlapping mechanisms: perceived functional legibility, visual consistency across contexts, adaptation burden, and the cognitive or interactional costs of switching between multiple device forms[24],[27],[28].

In this study, the form-based role cue was operationalized as a contrast between task-specific

embedded forms and a consistent standalone smart-speaker form. This manipulation was initially designed to approximate a Specialist-Generalist distinction. However, because this contrast also varies the degree of form consistency across tasks, any observed differences should not be interpreted as pure evidence of role specialization.

RQ2 exploratory: In routine task contexts, how are task-specific embedded forms versus a consistent standalone form associated with users' perceived ease of use, perceived usefulness, trust, attitude, and intention to use?

2-4 TAM as a Mediating Framework

The Technology Acceptance Model (TAM) posits that perceived ease of use (PEOU) and perceived usefulness (PU) are central determinants of technology acceptance[13],[14],[29]. In the present study, TAM serves as the mediating framework through which visual anthropomorphic cues are expected to shape downstream evaluations. The key analytical question is whether cue effects are transmitted through usability perceptions rather than directly.

The study examines three outcomes: trust, attitude, and intention to use. These are treated as parallel outcomes rather than a causal chain because the study captures pre-adoption evaluations formed after brief exposure to a simulated interaction[5],[7]. This approach allows the analysis to test whether PEOU and PU mediate the effects of social cues on each evaluative dimension while avoiding stronger causal assumptions among the outcomes themselves.

2-5 Hypotheses

H1: Perceived ease of use will be positively associated with (a) trust, (b) attitude, and (c) intention to use.

H2: Perceived usefulness will be positively associated with (a) trust, (b) attitude, and (c) intention to use.

H3: Perceived ease of use will be positively associated with perceived usefulness, and the serial indirect pathway from the visual anthropomorphic cue to outcomes through PEOU and PU will account for differences in (a) trust, (b) attitude, and (c) intention to use.

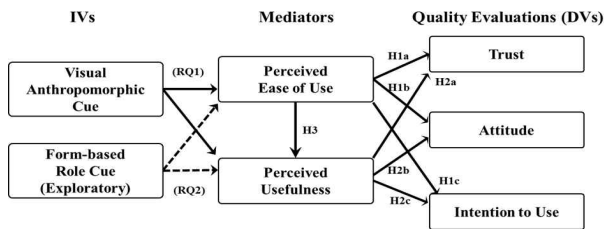


Fig. 1. Research model of the study. Solid lines indicate the primary TAM-based mediation model for the visual anthropomorphic cue. Dashed lines indicate the exploratory device-form contrast, initially designed as a form-based role cue

III. Method

3-1 Participants and Design

This study employed a 2 (Visual Cue: Human-like vs. Machine-like)×2 (Form-Based Role Cue: Specialist vs. Generalist) between-subjects online experiment. Participants were recruited in the Republic of Korea through Macromill Embrain and were eligible if they were 20–49 years old and had prior experience with voice-based AI assistants. The final analytic sample consisted of 195 participants (88 males, 107 females; $M_{age}=33.18$, $SD=8.08$), who were randomly assigned to one of the four conditions. Participants provided informed consent and received monetary compensation upon completion.

3-2 Experimental Stimuli and Procedure

Participants viewed video-based simulations of an AI conversational agent performing 12 routine tasks across bedroom, living-room, and vehicle settings. To minimize brand bias, the agent was presented under a neutral name (“UP”) without association to any commercial product. After viewing the assigned condition, participants completed the questionnaire.

The visual cue was manipulated through the agent’s screen feedback: the Human-like condition displayed animated eyes, whereas the Machine-like condition displayed speech data visualization. The form-based role cue was manipulated through device form: In the condition initially labelled Specialist, the agent was embedded into task-specific objects across tasks; in the condition initially labelled Generalist, it retained a

consistent standalone smart-speaker form. Because this manipulation necessarily varied both the intended role cue and the degree of form consistency across tasks, the resulting contrast is interpreted cautiously as a device-form contrast rather than as a pure manipulation of perceived expertise.

Table 1. Task assignments across settings

Setting	Tasks
Bedroom	Greeting, Checking weather, Controlling lights, Checking the time
Living room	Setting an alarm, Setting a timer, Controlling the TV, Ordering food
Vehicle	Checking vehicle condition, Navigation, Making a call, Sending a text message

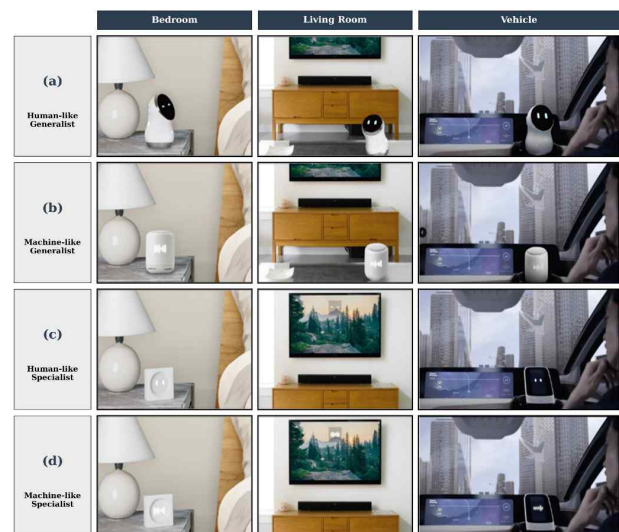


Fig. 2. Four experimental conditions of a conversational agent in three settings

3-3 Measurements

All constructs were measured on 5-point Likert scales using multi-item measures adapted from prior research. Perceived Ease of Use ($\alpha=.86$) and Perceived Usefulness ($\alpha=.88$) were adapted from Venkatesh[30] and Park et al.[31]. Trust ($\alpha=.84$), Attitude ($\alpha=.87$), and Intention to Use ($\alpha=.89$) were adapted from established prior studies[32]–[34]. Control Motivation ($\alpha=.80$) was included as a control variable alongside age and gender [35]. To improve measurement transparency, representative items for each construct are summarized in Table 2.

Table 2. Measurement constructs and representative items

Constructs	Representative item	α
Perceived Ease of Use	Using this AI conversational agent would be easy for me.	.86
Perceived Usefulness	This AI conversational agent would be useful for routine tasks.	.88
Trust	This AI conversational agent appears reliable.	.84
Attitude	Using this AI conversational agent would be a good idea.	.87
Intention	I would intend to use this AI conversational agent.	.89
Control Motivation	I prefer to maintain control when using automated systems.	.80

3-4 Manipulation Check

Manipulation checks were conducted to verify that participants perceived the experimental conditions as intended. Independent-samples *t*-tests were performed on dedicated manipulation check items administered after exposure to the stimuli.

Visual anthropomorphic cue. The visual cue manipulation was successfully differentiated. Participants in the Human-like condition rated the agent as significantly more human-like than those in the Machine-like condition (Q12: $M=2.29$, $SD=0.94$ vs. $M=1.75$, $SD=0.92$). Conversely, participants in the Machine-like condition rated the agent as significantly more machine-like (Q13: $M=4.06$, $SD=0.86$ vs. $M=3.76$, $SD=0.78$). These results confirm that the visual anthropomorphism manipulation was effective, with participants in the two conditions perceiving the agent's appearance in the intended directions at a medium effect size.

Form-based role cue. The form-based role cue manipulation yielded weaker differentiation. Neither the function-inferability item (Q17: $M_{\text{Spec}}=3.02$, $SD=0.97$ vs. $M_{\text{Gen}}=2.86$, $SD=1.01$) nor the function-expressing item (Q15: $M_{\text{Spec}}=3.52$, $SD=0.88$ vs. $M_{\text{Gen}}=3.49$, $SD=0.86$) reached statistical significance. This result indicates that the form-follows-function design, while conceptually distinct from the generalist form, was not reliably perceived by participants as signaling domain expertise or functional specialization in the manner intended.

This outcome is consistent with the conceptual distinction drawn in Section 2.3: form-based role cues may not straightforwardly activate the expertise heuristic. The specialist stimuli may have been perceived as functionally distinct objects (a lamp, a clock) rather than as domain-expert agents. Because

the form-based role cue did not achieve clear perceptual differentiation, all findings related to this variable are treated as exploratory throughout this paper and should not be interpreted as confirmatory evidence for the expertise heuristic in conversational agent design.

IV. Results

All analyses reported in this section are based on the final analytic sample of $N=195$ participants. Two-way analysis of variance (ANOVAs) were performed to identify main effects of the two experimental contrasts, followed by mediation and serial mediation analyses using the PROCESS macro (Models 4 and 6)[36] with 5,000 bootstrap resamples. The results are organized hierarchically: the primary visual anthropomorphic cue findings are reported first, followed by the divergent trust finding, and then the exploratory form-based role cue results.

4-1 Main Effects of Social Cues: ANOVA Results

Visual anthropomorphic cue (RQ1): The Machine-like condition was associated with significantly higher perceived ease of use than the Human-like condition ($M=3.99$, $SD=0.52$ vs. $M=3.73$, $SD=0.63$; $F(1, 191)=10.97$, $p<.01$, $\eta^2=.054$). The Machine-like condition was also associated with higher perceived usefulness ($F(1, 191)=4.02$, $p<.05$, $\eta^2=.021$), more favorable attitudes ($F(1, 191)=3.91$, $p<.05$, $\eta^2=.020$), and higher intention to use ($F(1, 191)=5.87$, $p<.05$, $\eta^2=.030$). The difference in trust between the two visual conditions was not statistically significant ($F(1, 191)=0.13$, $p>.05$, $\eta^2=.001$). In summary, human-like visual cues were associated with lower scores across all TAM and

behavioral outcome variables except trust, with the largest effect observed for perceived ease of use.

Form-based role cue (RQ2, exploratory): Because the manipulation check did not confirm the intended Specialist-Generalist distinction, these results are reported as an exploratory device-form contrast. The consistent standalone-form condition was associated with higher perceived ease of use than the task-specific embedded-form condition ($F(1, 191)=4.85, p<.05, \eta^2=.025$), as well as higher trust ($F(1, 191)=4.01, p<.05, \eta^2=.021$) and intention to use ($F(1, 191)=5.09, p<.05, \eta^2=.026$). Differences in perceived usefulness ($F(1, 191)=0.37, p>.05$) and attitude ($F(1, 191)=1.53, p>.05$) were not statistically significant. These patterns should be interpreted as possible responses to form consistency or adaptation burden, not as confirmed role-specialization effects.

No statistically significant interaction between the visual cue and the form-based role cue was observed for any dependent variable (all $ps>.05$). Given the exploratory status of the form-based contrast, this absence of interaction should be interpreted cautiously.

4-2 Mediation Analysis: Visual Anthropomorphic Cue

To examine whether the effect of the visual anthropomorphic cue on downstream outcomes was mediated through PEOU and PU, serial mediation analysis was conducted using PROCESS Model 6 [36] with 5,000 bootstrap resamples. In this analysis, the visual cue served as the independent variable (coded: 0=Machine-like, 1=Human-like), PEOU and PU served as serial mediators, and trust, attitude, and intention to use served as separate outcome variables. Age, gender, and control motivation were included as

covariates.

The Human-like visual cue was associated with significantly lower PEOU ($\beta=-0.27, p=.001$). This indicates that, in the present routine task context, participants who viewed the anthropomorphic agent perceived interaction as requiring more effort than those who viewed the machine-like agent.

PEOU was strongly and positively associated with PU ($\beta=0.58, p<.001$), confirming the well-established TAM pathway in which ease-of-use perceptions contribute to perceived usefulness. PEOU also directly predicted attitude ($\beta=0.18, p<.01$) and intention to use ($\beta=0.17, p<.05$), supporting H1b and H1c. PU in turn significantly predicted attitude ($\beta=0.63, p<.001$) and intention to use ($\beta=0.63, p<.001$), supporting H2b and H2c. The serial indirect effect (Visual Cue $\rightarrow$ PEOU $\rightarrow$ PU $\rightarrow$ Outcome) was significant and negative for both attitude (Effect=-0.10, SE=0.03, 95% CI [-0.17, -0.04]) and intention to use (Effect=-0.10, SE=0.04, 95% CI [-0.17, -0.04]), supporting H3b and H3c. After accounting for the mediators, the direct effect of the visual cue on attitude and intention were non-significant ($ps>.30$).

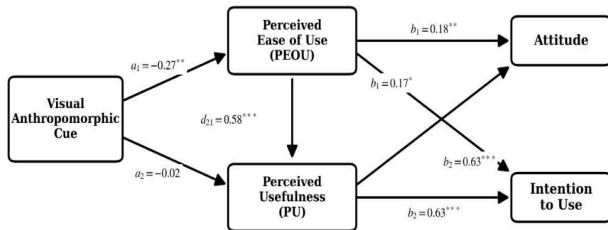
Neither the direct path from PEOU to trust nor the serial indirect effect through PU reached statistical significance in the visual cue model ($p>.05$). H1a and H2a were therefore not supported. This finding indicates that the mechanism by which visual anthropomorphic cues affected attitude and intention to use—through reduced ease of use and subsequent lower perceived usefulness—did not extend to trust. Trust, as a cognitive assessment of reliability, may depend on additional antecedents (e.g., perceived competence, transparency, or consistency of agent behavior) that were not differentially activated by the

Table 3. The main effects of social cues

DVs	Machine-like M(SD)	Human-like M(SD)	F (η^2)	Generalist M(SD)	Specialist M(SD)	F (η^2)
Ease of Use	3.993 (.522)	3.726 (.623)	10.97** (.054)	3.944 (.560)	3.855 (.590)	4.85* (.025)
Usefulness	3.738 (.628)	3.561 (.638)	4.02* (.021)	3.670 (.650)	3.646 (.637)	0.37 (.002)
Trust	3.218 (.594)	3.191 (.624)	0.13 (.001)	3.292 (.556)	3.204 (.608)	4.01* (.021)
Attitude	3.631 (.573)	3.465 (.631)	3.91* (.020)	3.597 (.575)	3.545 (.608)	1.53 (.008)
Intention	3.596 (.600)	3.370 (.732)	5.87* (.030)	3.586 (.645)	3.479 (.679)	5.09* (.026)

Note. * $p<.05$, ** $p<.01$

visual cue in this experiment. This divergence is addressed in Section 5.



Note. * $p < .05$, ** $p < .01$, *** $p < .001$

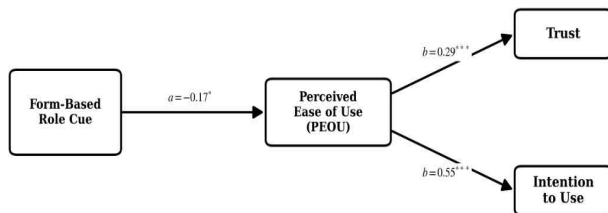
Fig. 3. Mediation effects of visual anthropomorphic cue

Table 4. Summary of key serial mediation paths

Outcome	Key paths	Serial indirect effect
Attitude	Visual cue $\rightarrow$ PEOU: -0.27^{**} ; PEOU $\rightarrow$ PU: 0.58^{***} ; PEOU $\rightarrow$ Attitude: 0.18^{**} ; PU $\rightarrow$ Attitude: 0.63^{***}	Effect = -0.10^* , 95% CI [$-0.17, -0.04$]
Intention	Visual cue $\rightarrow$ PEOU: -0.27^{**} ; PEOU $\rightarrow$ PU: 0.58^{***} ; PEOU $\rightarrow$ Intention: 0.17^* ; PU $\rightarrow$ Intention: 0.63^{***}	Effect = -0.10^* , 95% CI [$-0.17, -0.04$]

4-3 Exploratory Analysis: Form-Based Role Cue

An exploratory mediation analysis was conducted for the device-form contrast, rather than for a validated Specialist-Generalist distinction. The task-specific embedded-form condition was associated with lower PEOU than the consistent standalone-form condition ($\beta = -0.17$, $p = .039$). PEOU was positively associated with trust ($\beta = 0.30$, $p < .001$) and intention to use ($\beta = 0.55$, $p < .001$). These findings are suggestive only and may reflect form consistency or adaptation burden rather than role-type preference.



Note. * $p < .05$, *** $p < .001$.

Fig. 4. Exploratory mediation results for the form-based role cue

V. Discussion

This study found that human-like visual cues reduced PEOU in routine task contexts, and this effect propagated through PU to lower attitude and intention to use. Trust did not follow the same mediated pathway. The form-based role cue findings are retained only as exploratory because the intended Specialist-Generalist distinction was not confirmed, and may instead reflect device-form consistency or adaptation burden.

5-1 Theoretical Implications

The most robust finding of this study is that human-like visual cues reduced perceived ease of use in a routine task environment, with the effect fully mediated through the PEOU $\rightarrow$ PU $\rightarrow$ Attitude/Intention pathway. This result supports the view that anthropomorphism is context-dependent rather than uniformly beneficial. While prior research has often shown positive effects of anthropomorphic cues in relational or service-oriented settings, the present findings suggest that in utilitarian task environments, such cues may raise expectations that the system cannot meet, thereby undermining initial usability perceptions. This interpretation is consistent with expectation disconfirmation theory and with recent work reporting increased cognitive load or reduced usability in task-oriented AI and IoT contexts[8]–[10].

A second notable finding is that trust did not follow the same mediated pathway as attitude and intention to use. This suggests that trust in AI conversational agents may require evidence beyond initial ease-of-use and usefulness judgments, such as perceived competence, reliability, transparency, predictability, and behavioral consistency. In this experiment, functional performance, response content, voice output, and processing speed were held constant across conditions, which may explain why the visual cue affected attitude and intention but not trust. Thus, the findings indicate a boundary of the TAM-based pathway: PEOU and PU explained attitude and intention more clearly than trust in a short, video-based pre-adoption evaluation.

The device-form findings remain exploratory. Given that the intended Specialist-Generalist distinction was

not confirmed by the manipulation check, the observed differences are interpreted as possible responses to form consistency and adaptation burden rather than as expertise-based specialization effects. This account remains tentative and requires confirmatory testing.

5-2 Practical Implications

Within the limited context of short video-based pre-adoption evaluations, these findings cautiously suggest that visual design for routine-use AI conversational agents may benefit from prioritizing functional clarity over character-based expressiveness. Minimally expressive feedback may be more appropriate than human-like animated expressions when the primary goal is efficient task completion. The exploratory form-based role cue findings should be treated as provisional design considerations for future user experiment rather than direct prescriptions for product design.

5-3 Limitations and Future Research

This study has several limitations. First, the anthropomorphic manipulation was limited to a low-fidelity visual cue displayed on a screen, and thus the findings should not be generalized to multimodal or high-fidelity anthropomorphism involving voice, gesture, or physical embodiment. Second, the form-based role cue did not achieve statistically significant perceptual differentiation. Therefore, all related findings must be regarded as exploratory and should not be interpreted as evidence of a confirmed Specialist-Generalist role effect. Moreover, this manipulation confounded the intended role cue with device-form consistency: the Specialist condition used multiple task-specific embedded forms, whereas the Generalist condition used a single consistent standalone form. Thus, the observed pattern may reflect lower adaptation burden or reduced cognitive cost in the consistent-form condition rather than perceived role generality itself. Future research should disentangle these mechanisms by independently manipulating role framing, form consistency, and demonstrated task competence, such as explicit expertise labels, behavioral demonstrations of specialization, or factorial designs that hold device form constant while varying role information. Finally,

given that the observed effects were small to moderate and obtained from a single-session, video-based pre-adoption evaluation, the findings should not be generalized to actual long-term use without further interactive validation.

VI. Conclusion

This study showed that, in routine task contexts, human-like visual cues on AI conversational agents were associated with lower perceived ease of use than machine-like feedback, and that this difference propagated through perceived usefulness to reduce attitude and intention to use. These findings suggest that in utilitarian settings, visual anthropomorphic cues may undermine initial usability perceptions by creating expectations that exceed what the system can deliver. Trust did not follow the same pathway, indicating that trust formation may depend on antecedents beyond usability perceptions. Taken together, the findings indicate that visual anthropomorphic cues should be calibrated to pragmatic task demands rather than assumed to be uniformly beneficial. The device-form findings remain exploratory and require validation in interactive or longitudinal settings.

Acknowledgement

This work was supported by the Ministry of Education of the Republic of Korea and the National Research Foundation of Korea (NRF-2025S1A6B5A02003910).

Reference

- [1] S. S. Sundar, "Rise of Machine Agency: A Framework for Studying the Psychology of Human - AI Interaction (HAI)," *Journal of Computer-Mediated Communication*, Vol. 25, No. 1, pp. 74-88, January 2020. <https://doi.org/10.1093/jcmc/zmz026>
- [2] S. S. Sundar, H. Jia, T. F. Waddell, and Y. Huang, Toward a Theory of Interactive Media Effects (TIME): Four Models for Explaining How Interface Features Affect User Psychology, in *The Handbook of the Psychology of*

- Communication Technology*, Chichester, UK: Wiley-Blackwell, pp. 47-86, 2015. <https://doi.org/10.1002/9781118426456.ch3>
- [3] B. Reeves and C. Nass, *The Media Equation: How People Treat Computers, Television, and New Media Like Real People and Places*, Cambridge, UK: Cambridge University Press, 1996.
- [4] N. Epley, A. Waytz, and J. T. Cacioppo, "On Seeing Human: A Three-Factor Theory of Anthropomorphism," *Psychological Review*, Vol. 114, No. 4, pp. 864-886, 2007. <https://doi.org/10.1037/0033-295X.114.4.864>
- [5] M. Blut, C. Wang, N. V. Wunderlich, and C. Brock, "Understanding Anthropomorphism in Service Provision: A Meta-Analysis of Physical Robots, Chatbots, and Other AI," *Journal of the Academy of Marketing Science*, Vol. 49, No. 4, pp. 632-658, 2021. <https://doi.org/10.1007/s11747-020-00762-y>
- [6] E. Konya-Baumbach, M. Biller, and S. von Janda, "Someone Out There? A Study on the Social Presence of Anthropomorphized Chatbots," *Computers in Human Behavior*, Vol. 139, 107513, February 2023. <https://doi.org/10.1016/j.chb.2022.107513>
- [7] S. Moussawi, M. Koufaris, and R. Benbunan-Fich, "How Perceptions of Intelligence and Anthropomorphism Affect Adoption of Personal Intelligent Agents," *Electronic Markets*, Vol. 31, No. 2, pp. 343-364, June 2021. <https://doi.org/10.1007/s12525-020-00411-w>
- [8] G. Z. Karimova, "Not in Our Image: Rethinking Anthropomorphism in Expert Chatbot Design," *AI & Society*, Vol. 41, pp. 611-628, July 2025. <https://doi.org/10.1007/s00146-025-02438-z>
- [9] H. Kang and K. J. Kim, "Does Humanization or Machinization Make the IoT Persuasive? The Effects of Source Orientation and Social Presence," *Computers in Human Behavior*, Vol. 129, 107152, April 2022. <https://doi.org/10.1016/j.chb.2021.107152>
- [10] A.-M. Seeger, J. Pfeiffer, and A. Heinzl, "Texting with Humanlike Conversational Agents: Designing for Anthropomorphism," *Journal of the Association for Information Systems*, Vol. 22, No. 4, pp. 931-967, July 2021. <https://doi.org/10.17705/1jais.00685>
- [11] Y. J. Koh and S. S. Sundar, "Heuristic Versus Systematic Processing of Specialist Versus Generalist Sources in Online Media," *Human Communication Research*, Vol. 36, No. 2, pp. 103-124, April 2010. <https://doi.org/10.1111/j.1468-2958.2010.01370.x>
- [12] A. Rapp, L. Curti, and A. Boldi, "The Human Side of Human-Chatbot Interaction: A Systematic Literature Review of Ten Years of Research on Text-Based Chatbots," *International Journal of Human-Computer Studies*, Vol. 151, 102630, July 2021. <https://doi.org/10.1016/j.ijhcs.2021.102630>
- [13] F. D. Davis, "Perceived Usefulness, Perceived Ease of Use, and User Acceptance of Information Technology," *MIS Quarterly*, Vol. 13, No. 3, pp. 319-340, September 1989. <https://doi.org/10.2307/249008>
- [14] V. Venkatesh and F. D. Davis, "A Theoretical Extension of the Technology Acceptance Model: Four Longitudinal Field Studies," *Management Science*, Vol. 46, No. 2, pp. 186-204, February 2000. <https://doi.org/10.1287/mnsc.46.2.186.11926>
- [15] K. M. Lee, "Presence, Explicated," *Communication Theory*, Vol. 14, No. 1, pp. 27-50, February 2004. <https://doi.org/10.1111/j.1468-2885.2004.tb00302.x>
- [16] L. Qiu and I. Benbasat, "Evaluating Anthropomorphic Product Recommendation Agents: A Social Relationship Perspective to Designing Information Systems," *Journal of Management Information Systems*, Vol. 25, No. 4, pp. 145-182, 2009. <https://doi.org/10.2753/MIS0742-1222250405>
- [17] Y. J. Koh and S. S. Sundar, "Effects of Specialization in Computers, Web Sites, and Web Agents on E-Commerce Trust," *International Journal of Human-Computer Studies*, Vol. 68, No. 12, pp. 899-912, December 2010. <https://doi.org/10.1016/j.ijhcs.2010.08.002>
- [18] R. L. Oliver, "A Cognitive Model of the Antecedents and Consequences of Satisfaction Decisions," *Journal of Marketing Research*, Vol. 17, No. 4, pp. 460-469, November 1980. <https://doi.org/10.2307/3150499>
- [19] R. Kocielnik, S. Amershi, and P. N. Bennett, "Will You Accept an Imperfect AI? Exploring Designs for Adjusting End-User Expectations of AI Systems," in *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems*, Glasgow, Scotland, pp. 1-14, May 2019. <https://doi.org/10.1145/3290605.3300641>
- [20] C. Nass and K. M. Lee, "Does Computer-Synthesized Speech Manifest Personality? Experimental Tests of Recognition, Similarity-Attraction, and Consistency-Attraction," *Journal of Experimental Psychology: Applied*, Vol. 7, No. 3, pp. 171-181, 2001. <https://doi.org/10.1037/1076-898X.7.3.171>
- [21] D. Kontogiorgos, G. Skantzé, A. Pereira, and J. Gustafson, "The Effects of Embodiment and Social Eye-Gaze in Conversational Agents," in *Proceedings of the 41st Annual Conference of the Cognitive Science Society (CogSci 2019)*, Montreal, Canada, July 2019. <https://cogsci.mindmo>

- deling.org/2019/papers/0119/index.html
- [22] C. D. Kidd and C. Breazeal, "Effect of a Robot on User Perceptions," in *Proceedings of the 2004 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, Sendai, Japan, Vol. 4, pp. 3559-3564, September 2004. <https://doi.org/10.1109/IROS.2004.1389967>
- [23] C. Nass, B. Reeves, and G. Leshner, "Technology and Roles: A Tale of Two TVs," *Journal of Communication*, Vol. 46, No. 2, pp. 121-128, June 1996. <https://doi.org/10.1111/j.1460-2466.1996.tb01477.x>
- [24] D. Norman, *The Design of Everyday Things*, Revised and Expanded ed. New York, NY: Basic Books, 2013.
- [25] H. S. Nwana, "Software Agents: An Overview," *The Knowledge Engineering Review*, Vol. 11, No. 3, pp. 205-244, September 1996. <https://doi.org/10.1017/S02698890000789X>
- [26] L. Sullivan, "The Tall Office Building Artistically Considered," *Lippincott's Magazine*, Vol. 57, pp. 403-409, March 1896.
- [27] D. L. Goodhue and R. L. Thompson, "Task-Technology Fit and Individual Performance," *MIS Quarterly*, Vol. 19, No. 2, pp. 213-236, June 1995. <https://doi.org/10.2307/249689>
- [28] K. J. Kim, "Interacting Socially with the Internet of Things (IoT): Effects of Source Attribution and Specialization in Human-IoT Interaction," *Journal of Computer-Mediated Communication*, Vol. 21, No. 6, pp. 420-435, November 2016. <https://doi.org/10.1111/jcc4.12177>
- [29] R. J. Holden and B.-T. Karsh, "The Technology Acceptance Model: Its Past and Its Future in Health Care," *Journal of Biomedical Informatics*, Vol. 43, No. 1, pp. 159-172, February 2010. <https://doi.org/10.1016/j.jbi.2009.07.002>
- [30] V. Venkatesh, "Determinants of Perceived Ease of Use: Integrating Control, Intrinsic Motivation, and Emotion into the Technology Acceptance Model," *Information Systems Research*, Vol. 11, No. 4, pp. 342-365, December 2000. <https://doi.org/10.1287/isre.11.4.342.11872>
- [31] N. Park, Y.-C. Kim, H. Y. Shon, and H. Shim, "Factors Influencing Smartphone Use and Dependency in South Korea," *Computers in Human Behavior*, Vol. 29, No. 4, pp. 1763-1770, July 2013. <https://doi.org/10.1016/j.chb.2013.02.008>
- [32] B. J. Fogg and H. Tseng, "The Elements of Computer Credibility," in *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, Pittsburgh, PA, pp. 80-87, May 1999. <https://doi.org/10.1145/302979.303001>
- [33] C. E. Porter and N. Donthu, "Using the Technology

Acceptance Model to Explain How Attitudes Determine Internet Usage: The Role of Perceived Access Barriers and Demographics," *Journal of Business Research*, Vol. 59, No. 9, pp. 999-1007, September 2006. <https://doi.org/10.1016/j.jbusres.2006.06.003>

- [34] C. M. Jackson, S. Chow, and R. A. Leitch, "Toward an Understanding of the Behavioral Intention to Use an Information System," *Decision Sciences*, Vol. 28, No. 2, pp. 357-389, April 1997. <https://doi.org/10.1111/j.1540-5915.1997.tb01315.x>
- [35] N. Epley, A. Waytz, S. Akalis, and J. T. Cacioppo, "When We Need a Human: Motivational Determinants of Anthropomorphism," *Social Cognition*, Vol. 26, No. 2, pp. 143-155, April 2008. <https://doi.org/10.1521/soco.2008.26.2.143>
- [36] A. F. Hayes, *Introduction to Mediation, Moderation, and Conditional Process Analysis: A Regression-Based Approach*, 3rd ed. New York, NY: Guilford Press, 2022.



정재열 (Jaeyoul Jeong)

2011년 : 연세대학교 (학사-패키징학)

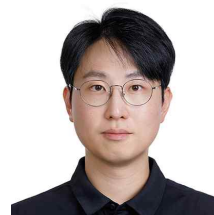
2019년 : 성균관대학교 (박사-인간
컴퓨터 상호작용, HCI)

2018년~2021년: 채단법인 스마트IT융합시스템연구단 연구교수

2019년~2021년: ㈜유벤처파트너스

2024년~2025년: 지자체-대학 협력기반 지역혁신 사업 대전 ·
세종 · 충남 지역혁신플랫폼 모빌리티ICT사
업본부 책임연구원

2025년~현 재: 한국과학기술정보연구원(KISTI) 선임기술원
※관심분야 : 인간-컴퓨터 상호작용(HCI), 사용자 경험(UX),
지능형 서비스, 데이터 거버넌스



윤홍석 (Hongsuk Yoon)

2013년 : 중앙대학교 (학사-철학,
사회학)

2017년 : 성균관대학교 (석사-인간
컴퓨터 상호작용, HCI)

2024년 : 유니버시티 칼리지 더블린
(박사-포용적 디자인,
Inclusive Design)

2015년~2017년: 인터랙션사이언스 소셜컴퓨팅랩 연구원

2017년~2024년: 아일랜드 스마트랩 연구원 (IDRC)

2025년~현 재: 건국대학교 모빌리티인문학 연구원 HK
연구교수

※관심분야 : 인간-컴퓨터 상호작용(HCI), 포용적 디자인
(Inclusive Design), 가상현실(VR)